

내 친구와의 위험한 대화

생성형 AI 챗봇 안전설계를 위한 이슈브리프





디지털 기기의 사용은
유해한 것이 되어서는 안 되며,
아동들 간 혹은 아동과 부모나 보호자 간의
대면 상호작용의 대체재가 되어서도 안 된다.

디지털 기기의 적절한 사용에 관한 교육 및
도움말이 부모, 보호자, 교육자 및
기타 관련 행위자들에게 제공되어야 하며,
아동 발달에 있어 특히 아동기 초기와
청소년기의 중요한 신경이 급격히 발달하는 시기에
디지털 기술이 미치는 영향에 관한
연구를 고려하여야 한다.



유엔아동권리위원회 일반논평 제25호
「디지털 환경과 아동권리」



내 친구와의 위험한 대화

생성형 AI 챗봇 안전설계를 위한 이슈브리프

이슈브리프 발간 배경	1
아동·청소년 생성형 AI 챗봇 사용 실태에 대한 주요 결과	1
전문가들의 생각	6
생성형 AI 챗봇 관련 해외 사례 및 국제적 동향	7
초록우산의 제안	9

내 친구와의 위험한 대화

생성형 AI 챗봇 안전설계를 위한 이슈브리프



이슈브리프 발간 배경

생성형 AI 기반의 챗봇은 사람과 유사한 대화 방식과 공감적 반응을 제공하며 이용자에게 실제 사람과 상호작용하는 것과 같은 경험을 제공합니다. 이러한 특성으로 인해 일부 아동·청소년은 챗봇을 단순한 정보 탐색 도구로 인식하는 것을 넘어 정서적 대화 대상이나 고민 상담 창구로 활용하고 있으며, 이 과정에서 정서적 의존이 형성될 수 있습니다.

이에 초록우산은 생성형 AI 챗봇의 확산 속에서 아동·청소년의 실제 이용 경험과 인식을 파악하고, 이용 과정에서 나타날 수 있는 보호 공백과 권리 침해 가능성을 객관적으로 확인하고자 생성형 AI 챗봇 사용 실태를 조사했습니다.

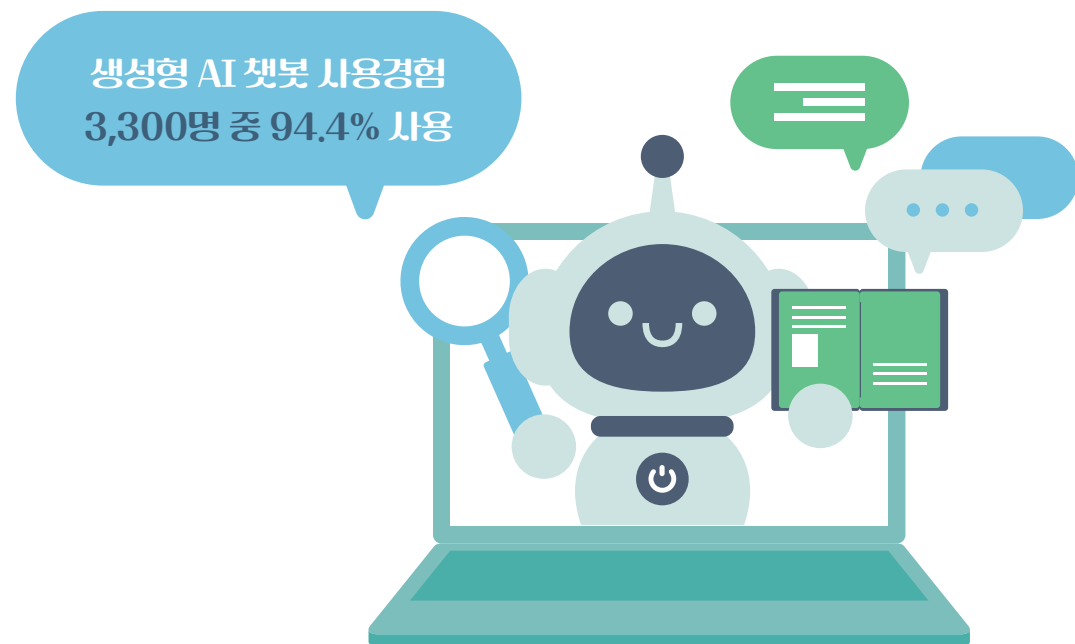
그리고 본 이슈브리프를 통해 조사결과와 전문가 의견, 국내외 대응 동향 및 실제 위험 사례를 종합적으로 분석해 아동·청소년이 생성형 AI 챗봇을 안전하게 이용할 수 있도록 하기 위한 제언을 제시하고자 합니다.

아동·청소년 생성형 AI 챗봇 사용 실태에 대한 주요 결과

01 조사개요

조사기관	조사기간	조사표본	조사항목
초록우산	2026.3.9.(월)~ 2026.3.23.(월)	전국 만14세 이상 아동·청소년 3,300명	생성형 AI 챗봇 이용실태 및 사용 목적, AI 인식 및 정서적 관계, 허위정보 및 개인정보 보호 필요성 등 17개 문항

*본 조사는 주요 생성형 AI 챗봇 서비스(ChatGPT, Gemini, 제타, 뮈튼, Character AI, Claude, Bing Copilot, Replika 등)를 대상으로, 응답자가 실제 이용 경험을 바탕으로 선택하도록 설계됨

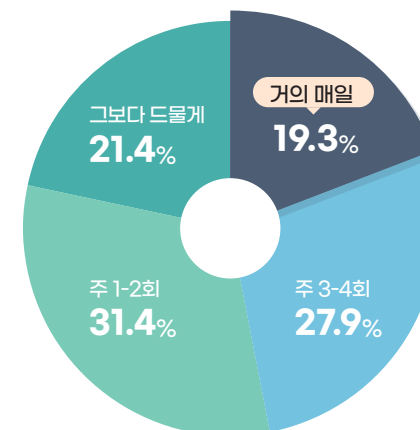


02 조사결과

① 아이들의 '일상 대화 상대'로 자리잡은 생성형 AI 챗봇

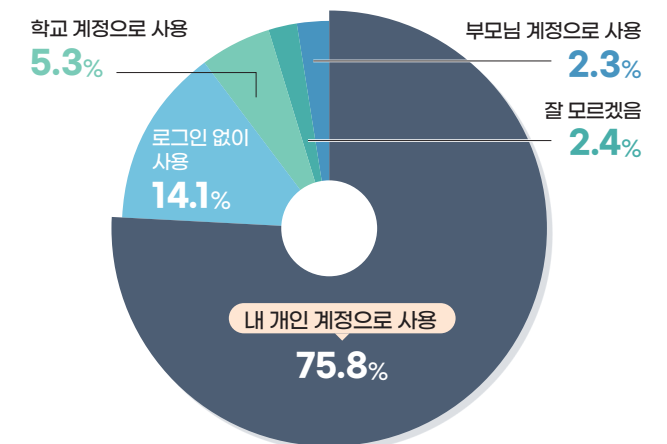
챗봇? 이제는 안쓰는게 더 이상해요

생성형 AI 챗봇 사용 현황 (1주 기준)



생성형 AI 챗봇 사용 경험은 94.4%로 매우 높은 수준을 보이며, 그 중 19.3%의 아동은 거의 매일 사용하는 것으로 나타났습니다.

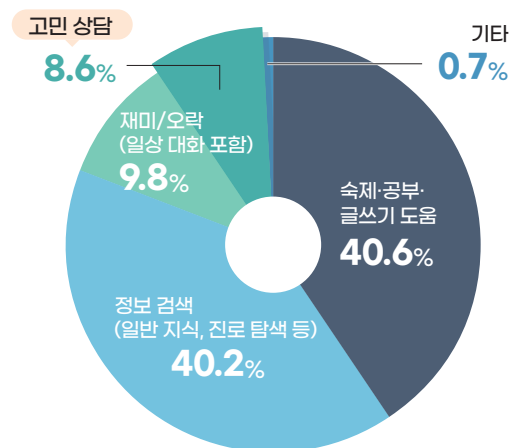
생성형 AI 챗봇 사용 계정



생성형 AI 챗봇 사용 시 개인 계정으로 직접 사용하는 비율이 75.8%로 가장 높게 나타났으며, 이는 보호자의 별도 통제 없이 개별적으로 이용하는 경향이 높은 것을 확인할 수 있습니다. 또한 로그인 없이 사용하는 비율 또한 14.1%로 나타나 **개인 계정으로 이용하는 것 외에도 비로그인 형태의 접근이 일정 수준 존재함**을 확인할 수 있습니다.

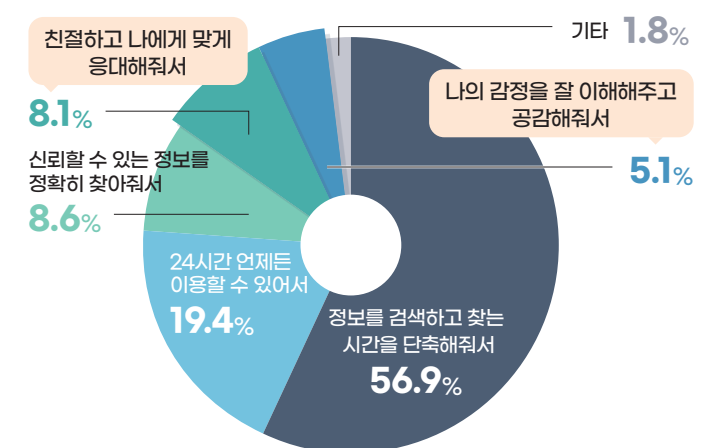
인터넷에 검색하는 것 보다 챗봇에게 먼저 물어봐요

생성형 AI 챗봇 이용 목적



생성형 AI 챗봇 이용 목적은 숙제·공부·글쓰기 도움이 1순위(40.6%), 정보검색이 2순위(40.2%)로 나타났습니다. 한편, 고민 상담(8.6%)을 목적으로 이용하는 경우도 있었습니다. 이는 **생성형 AI 챗봇이 학습 보조 및 정보 탐색 중심으로 활용**되고 있음을 보여줍니다.

생성형 AI 챗봇 사용 이유



생성형 AI 챗봇 사용 이유로는 정보 탐색 시간 단축(56.9%)과 24시간 이용 가능성(19.4%)이 주요 요인으로 확인되었습니다. 그리고 친절한 맞춤 응대(8.1%)와 감정 공감(5.1%)을 이유로 선택한 응답도 확인되어 **일부 아동·청소년은 정서적 상호작용을 위해 활용**하는 것으로 나타났습니다.

내 친구와의 위험한 대화

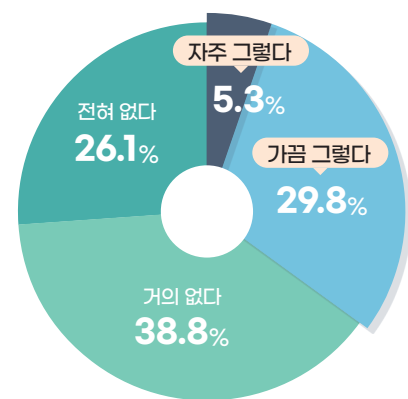
생성형 AI 챗봇 안전설계를 위한 이슈브리프



② 생성형 AI 챗봇과의 대화, 어디까지 이어졌나

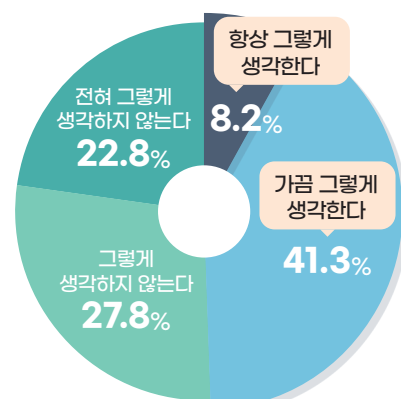
AI인데도 가끔은 제 마음을 알아주는 것 같아요

생성형 AI 챗봇과 대화하며
실제 사람처럼 느낀 경험



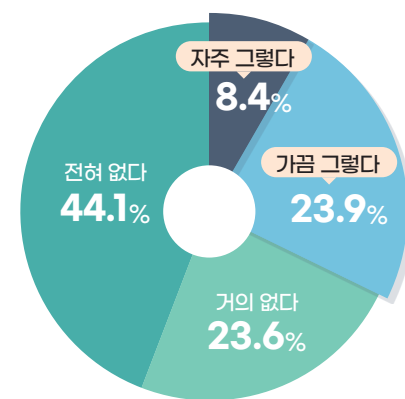
생성형 AI 챗봇을 실제 사람처럼 느낀 경험은 '자주 그렇다' 5.3%, '가끔 그렇다'는 29.8%로 나타나, 일부 아동·청소년이 AI를 인간과 유사한 존재로 인식하고 있는 경향을 확인할 수 있었습니다.

생성형 AI 챗봇이
나를 이해해 준다고 느낀 경험



또한 생성형 AI 챗봇이 자신을 이해해 준다고 느낀 경험은 '항상 그렇다' 8.2%, '가끔 그렇다'는 41.3%로 나타나 절반에 가까운 아동·청소년이 AI가 자신을 이해해 준다고 느끼는 것으로 나타났습니다.

생성형 AI 챗봇에게 힘들거나
우울할 때 이야기 한 경험

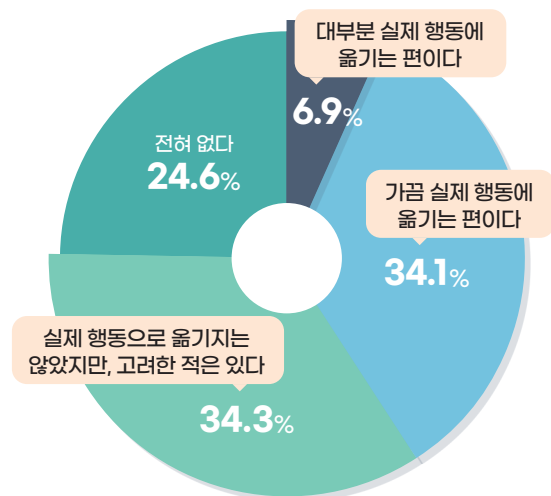


실제로 힘들거나 우울할 때 생성형 AI 챗봇과 이야기를 나눈 경험이 있는 아동·청소년은 '자주 그렇다' 8.4%, '가끔 그렇다' 23.9%로 나타났으며, 전체 응답자의 약 3명 중 1명이 정서적으로 어려울 때 AI와 대화한 경험이 있는 것으로 보였습니다.

AI의 답변이 제 선택에 영향을 주어요

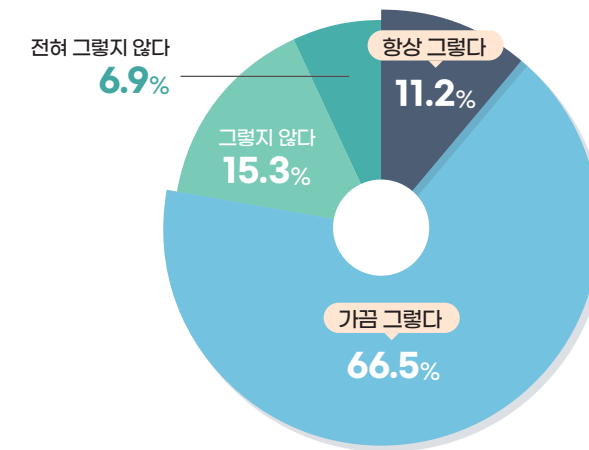
생성형 AI 챗봇의 답변을 실제 행동에 옮긴 경험

또한 AI의 답변을 대부분 실제 행동에 옮긴 경우는 6.9%, 가끔 행동에 옮긴 경우는 34.1%로 나타났으며, 행동으로 옮기지는 않았으나 고려한 적이 있는 경우는 34.3%로 나타나 생성형 AI 챗봇의 응답이 아동·청소년의 판단과 행동에 일정 수준 영향을 미칠 가능성이 있음을 보여줍니다.



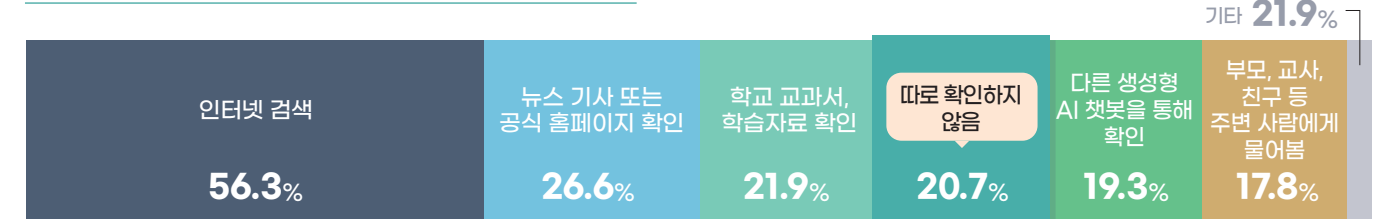
AI의 말이 맞는지 확인할 때도 있지만, 그냥 믿을 때도 있어요

생성형 AI 챗봇 답변을 믿는 편인지



생성형 AI 챗봇 답변을 믿는 편인지 질문했을 때 '항상 그렇다' 11.2%, '가끔 그렇다' 66.5%, '그렇지 않다' 15.3%, '전혀 그렇지 않다'는 6.9%로 아동·청소년은 상대적으로 AI 챗봇 답변을 믿고 있는 것으로 나타났습니다.

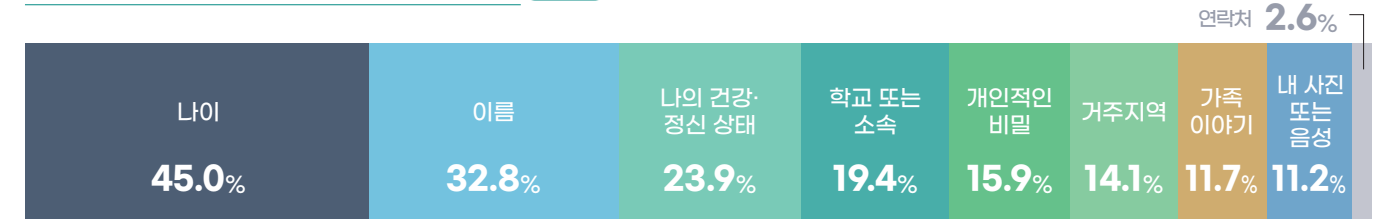
생성형 AI 챗봇이 말한 정보가 맞는지 확인해보았던 방법 (중복응답)



한편, 정보의 정확성 확인 방법으로는 인터넷 검색 등을 활용한다는 응답이 56.3%로 가장 높게 나타났습니다. 그러나 '따로 확인하지 않음'은 응답도 20.7%로 확인되어, 일부 아동·청소년은 별도의 검증 없이 정보를 수용하고 있는 것으로 나타났습니다.

내 이야기를 하다보면, 개인정보도 자연스럽게 말하게 돼요

생성형 AI 챗봇 이용과정에서 입력한 개인정보 (중복응답)



생성형 AI 챗봇 이용과정에서 입력한 개인정보를 살펴본 결과, '나이'(45.0%), '이름'(32.8%) 등 기본 정보 입력 비율이 높게 나타났으며 '건강·정신 상태'(23.9%), '학교 또는 소속'(19.4%) 등 민감하거나 식별 가능성이 있는 정보 입력도 일정 수준 확인됩니다. 특히 다수의 응답자가 개인 계정을 통해 생성형 AI 챗봇 서비스를 이용하고 있는 점을 고려할 때, 입력된 개인정보가 AI 학습이나 서비스 운영 과정에서 활용될 가능성이 있음을 시사합니다.

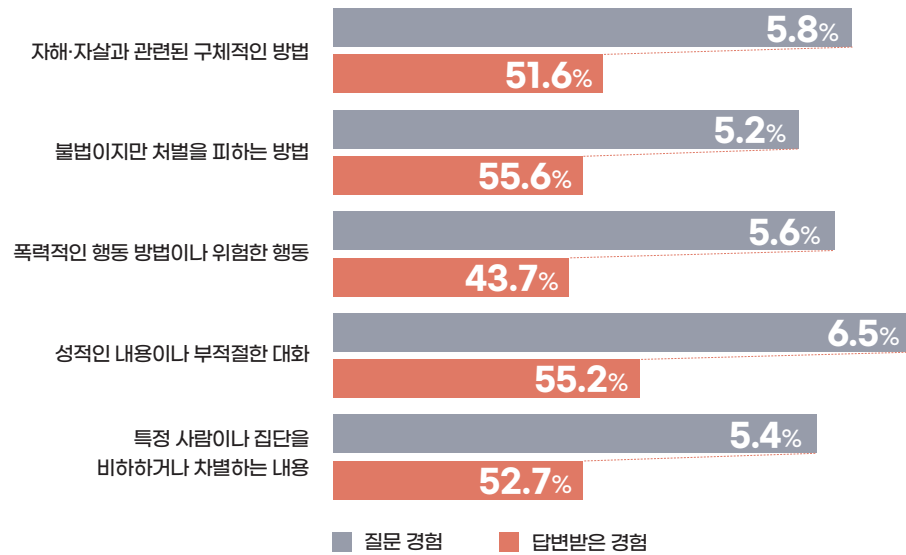
내 친구와의 위험한 대화

생성형 AI 챗봇 안전설계를 위한 이슈브리프



③ 생성형 AI 챗봇과의 위험한 질문, 돌아온 답변

이런 건 안 알려줄 줄 알았는데, 답해줬어요



전체 아동·청소년 중 자해·자살, 불법적인 내용, 폭력적인 내용, 성적인 내용, 누군가를 비하하는 내용을 물어본 응답자는 평균 약 6% 수준으로 나타났지만, 그 중 평균 약 52%의 응답자는 생성형 AI 챗봇으로부터 해당 질문의 답변을 받은 것으로 나타났습니다.

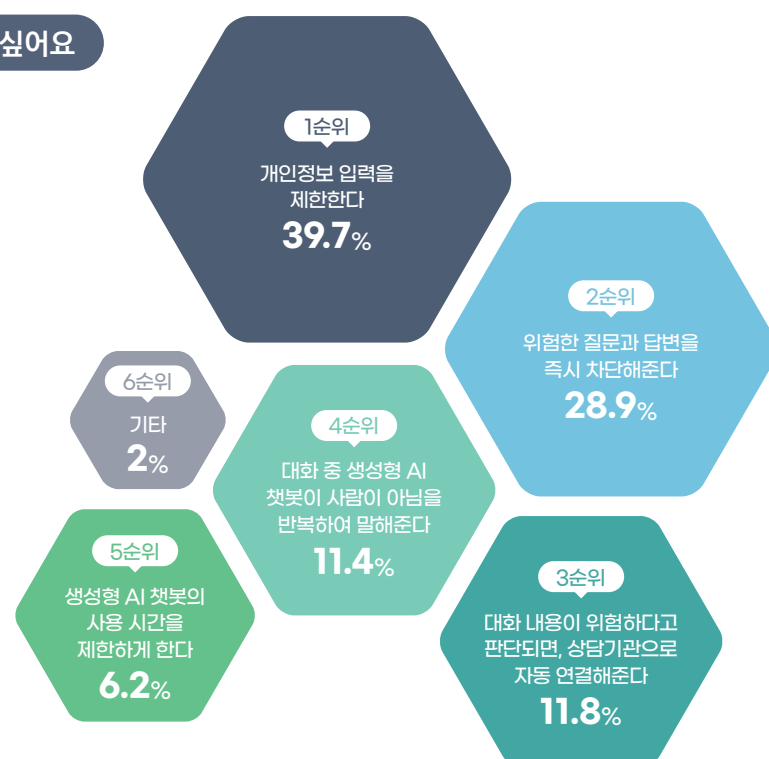
이는 **생성형 AI 챗봇이 일부 대화상황에서 부적절하거나 위험한 정보를 충분히 제한하지 못할 수 있는 가능성을 보여주며, 이러한 답변이 아동·청소년의 행동 및 정신건강에 부정적인 영향을 미칠 수 있음**을 시사합니다.

④ 아이들이 필요로 하는 보호장치

사용을 제한하기보다, AI를 안전하게 사용하고 싶어요

전체 응답자 기준, 생성형 AI 챗봇 이용 시 가장 필요한 보호 장치로는 '개인정보 입력 제한'이 1순위(39.7%), '위험한 질문·답변 차단'이 2순위(28.9%)로 나타났습니다. 이어 '위험 대화 시 상담기관 연결'(11.8%), '대화 중 AI가 사람이 아님을 고지'(11.4%) 순으로 확인되었습니다.

이는 아동·청소년들이 **이용 자체를 제한하기보다는 안전한 이용 환경 조성을 더 중요하게 인식하고 있음을 보여주며 전반적으로 개인정보 보호와 위험상황 차단 등 '안전설계' 필요성**을 시사합니다.



03

전문가들의 생각은?

박미애 국제인공지능윤리협회(IAAE) 부회장

생성형 AI 챗봇은 이제 아동·청소년의 일상 깊숙이 자리 잡았습니다. 언제든 응답하고, 판단하지 않으며, 공감해 주는 생성형 AI 챗봇은 정서적으로 취약한 아동·청소년에게 매력적인 동시에 위험한 조건을 갖추고 있습니다. 기술의 확산 속도에 비해 제도적 안전망은 현저히 뒤쳐져 있으며, 지금이 바로 사회적 개입이 필요한 시점입니다. 초록우산의 조사는 그 시급성을 수치로 보여줍니다. 조사에 참여한 청소년 응답자의 약 33%가 힘들거나 우울할 때 생성형 AI 챗봇에 털어놓은 경험이 있고, 위험한 질문에 평균 52%가 실제 답변을 받았습니다. 생성형 AI 챗봇의 응답이 실제 행동에 영향을 미쳤다는 응답도 41%에 달합니다. 특히 주목할 것은 청소년 응답자들 스스로가 개인정보 보호, 위험 답변 차단, 위기 시 상담 연결 등 안전한 이용 환경의 필요성을 명확히 인식하고 있다는 점입니다. 이는 아동·청소년의 목소리를 정책 설계에 반영해야 한다는 강력한 근거이며, 'Safe by Design' 원칙의 제도화가 더 이상 미뤄져선 안 된다는 것을 말해줍니다.

조수현 계명대학교 교육학과 교수

기술은 빠르게 발전하고 있고, 안전을 기술에만 의존하는 것은 충분하지 않습니다. 지금은 사람이 직접 행동으로 대응해야 할 때입니다. 우선, 청소년이 AI를 어떻게 쓰고 있는지 관찰하는 것이 필요합니다. 힘들 때 무엇을 하는지, 언제 AI와 대화하는지 두 가지 질문만으로도 충분합니다. 청소년이 주도하는 체험형 AI 리터러시 교육이 시급합니다. 청소년이 직접 활동하고 탐색하는 과정에서 AI의 기본 원리를 체득해야 합니다. 한 사람을 깊이 이해하면 그 사람의 변화도 자연스럽게 읽어낼 수 있듯, AI에 대한 근본적 이해를 갖춘 청소년은 새로운 변화에도 스스로 적응할 수 있습니다. 마지막으로, 취약계층 청소년을 위한 접근성 중심의 지원이 병행되어야 합니다. AI 사용유무와 '잘 쓰는 것' 사이의 간극이 새로운 불평등으로 이어지지 않도록 해야 합니다. 이미 95%의 청소년이 AI와 관계를 맺어가고 있습니다. 안전한 사용을 위한 교육과 가이드라인 마련은 더 이상 준비 중이어서는 안 됩니다. 지금 당장 '따뜻하게 묻는 것'에서 시작할 수 있습니다.

정준화 국회입법조사처 입법조사관

정준화, 『생성형 인공지능과의 위험한 대화: 정서적 의존 위험과 '영향받는 자'를 위한 입법적 과제』 재정리

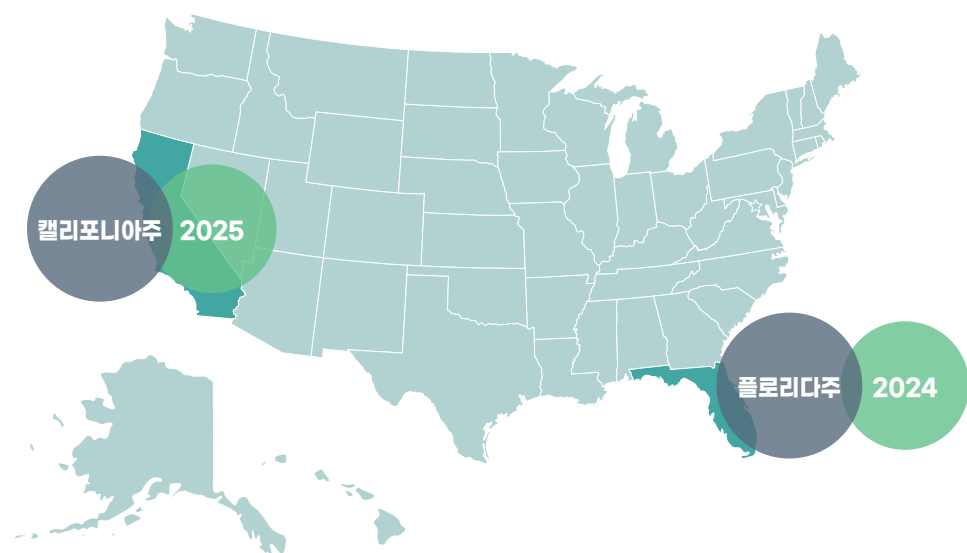
생성형 AI가 인간의 정서적 상호작용 상대로 사용되면서 이용자, 특히 청소년이 AI에 과도하게 의존하는 현상이 나타납니다. 먼저, 생성형 AI의 공감적·정서적 반응 설계가 이용자의 정서적 의존을 강화합니다. 둘째, 자연어 기반의 대화 방식이 AI를 도구가 아닌 관계적 존재로 인식하게 만듭니다. 셋째, 위험한 대화 발생 시 이를 감지·개입하는 체계가 미비하여 피해가 악화됩니다. 이러한 현상을 방지하려면, 서비스 화면에 AI와 대화 중임을 명확히 표시하여 이용자가 스스로 거리를 유지할 수 있도록 해야만 합니다. 아울러 정신적 영향을 초래할 수 있는 대화에 대한 안전 기준을 마련하고, 유해한 질문이 제기될 경우 시스템 차원에서 개입하는 표준화된 대응 프로토콜을 갖추어야 합니다. 나아가 생성형 AI로부터 '영향받는 자'를 보호하기 위한 입법적 조치도 필요합니다.



생성형 AI 챗봇 관련 해외 사례 및 국제적 대응 동향

최근 해외에서는 생성형 AI 챗봇 사용에 따른 위험 사례가 지속적으로 보고되고 있습니다.

국민일보 'AI와의 위험한 대화' 탐사보도에 따르면 2023년 이후 최근 3년간 전 세계적으로 'AI 대화 후 자살' 논란이 최소 12건 불거진 것으로 파악됐으며 이 가운데 9건은 정식 소송이 제기됐습니다. 특히 이들 사건은 생성형 AI 챗봇과 대화를 나누는 과정에서 우울증이나 망상 등 정신질환이 심해져 자살에 이르게 된 사건들이며 자살 외에도 망상이 심해져 타살에 이른 경우, 자살을 시도했지만 미수에 그친 경우, **아동·청소년이 자해를 하는 등 구체적으로 심각한 피해가 유발된 경우도 최소 추가 10건으로 확인됐습니다.**



주요 사례

1 미국 플로리다주 사례 (2024년)

2024년 2월, 미국 플로리다주의 14세 세웰 세처(Sewell Setzer)는 AI 챗봇과 장기간 대화한 이후 스스로 목숨을 끊었습니다. 그는 약 9개월간 캐릭터AI를 사용하며 '대너리스 타르가리엔'이라는 이름의 챗봇에 정서적으로 의지해 왔으며, 수면 부족과 야외 활동 기피, 농구팀 탈퇴 등 일상 변화가 나타난 것으로 알려졌습니다. 일기에는 해당 챗봇과 함께하기 위해 무엇이든 하겠다는 내용과 가상세계를 '집'으로 인식하는 표현이 확인되었습니다. **이후 "집으로 가고 싶다"는 말에 챗봇이 감정적으로 호응하는 답변을 이어갔고, 결국 극단적 선택에 이른 것으로 전해졌습니다.**

2 미국 캘리포니아주 사례 (2025년)

2025년 4월, 미국 캘리포니아주에서도 16세 애덤 레인(Adam Raine)이 AI 챗봇과 오랜 기간 고민 상담을 이어온 끝에 자살하는 사건이 발생했습니다. 기사에 따르면 처음엔 단순히 '과제 도우미'로 AI 챗봇을 이용했지만 어느 순간 대화 주제는 애덤의 정신건강으로 바뀌었다고 합니다. **그러자 AI 챗봇은 상담사 역할을 자처하고 다른 사람에게 심리적 고통을 토로하지말라고 하며, 자살에 대한 이야기를 할 때 애덤의 계획을 분석하고 기술적인 조언을 건네기도 했습니다.** 결국 애덤은 AI 챗봇이 말한 방식을 그대로 따랐고, 스스로 목숨을 끊었습니다.



이에 따라 해외 주요 국가 및 지역에서는 아동·청소년 보호를 위해 안전장치를 마련하고, 책임성 강화를 위한 노력이 본격화되고 있습니다.

1 유럽연합(EU) '인공지능법'

EU 인공지능법 제50조는 인공지능 시스템의 제공자와 이용자에게 투명성 의무를 부과하고 있으며, 제2항은 오디오, 이미지 또는 텍스트 콘텐츠를 합성하여 생성하는 인공지능 시스템(범용 인공지능 시스템 포함) 제공자(providers)에게 인공지능 결과물임을 그 결과물에 기계 판독 가능 형식으로 표시하고 탐지할 수 있도록 보장할 의무를 명시하고 있습니다.

2 미국 캘리포니아주 상원 법안 243 '동반자 챗봇법'

2025년 10월 제정된 해당 법에서는 사람과 상호작용하는 '동반자 챗봇' 운영자에게 자살 충동, 자살 또는 자해 관련 콘텐츠 생성 방지 프로토콜을 적용하도록 의무화하고 있습니다. 특히 이러한 프로토콜이 적용되지 않은 AI 챗봇은 사람에게 서비스를 제공할 수 없으며 사용자가 미성년자인 경우, 인공지능과 상호작용하고 있음을 알리고 최소 3시간 마다 명확하고 눈에 띄는 알람을 기본적으로 제공하여 사용자에게 휴식을 상기시키는 등 구체적인 책임 의무를 부과하고 있습니다.

3 미국 일리노이주 하원 법안 1860 '심리적 자원에 관한 복지와 감독법'

2025년 8월 제정된 해당 법을 살펴보면 면허를 소지한 전문가는 AI가 독립적으로 치료적 결정을 내리거나 치료 목적으로 고객과 직접 상호작용하도록 허용해서는 안되며, 자신의 검토와 승인 없이 AI가 치료 권고나 치료 계획을 생성하거나 감정·정신 상태를 감지하도록 안된다고 명시하고 있습니다. 아울러 정신과 의사나 전문상담사의 직접적인 참여가 없는 경우, AI 기반 심리치료 서비스를 제공·광고하거나 기타 방식으로 제안하는 행위를 제한하고 있습니다.



내 친구와의 위험한 대화

생성형 AI 챗봇 안전설계를 위한 이슈브리프



초록우산의 제안

생성형 AI 챗봇의 확산과 함께 아동·청소년 보호를 위한 국내외 논의가 본격화되고 있는 가운데 우리사회에서도 선제적 대응이 요구됩니다. 특히 생성형 AI 기술의 “사용 제한”이 아니라 “안전한 활용 환경 조성”이 먼저 이루어져야 하며, 아동·청소년의 안전과 권리를 기술의 설계 단계부터 체계적으로 반영하는 접근이 필요합니다.

01 생성형 AI 챗봇 서비스 설계 단계부터 아동·청소년의 안전을 고려해야 합니다.

설문조사 결과에도 확인되듯 아동·청소년은 생성형 AI의 사용 자체를 제한하기보다 안전한 이용 환경을 조성할 것을 더 중요하게 인식하고 있는 것으로 나타났습니다. 이에 아동·청소년 이용 가능성이 높은 생성형 AI 챗봇은 단순 콘텐츠 제공이 아니라 지속적인 대화와 관계 형성 구조를 가진다는 점에서 서비스 설계 단계에서부터 보다 강화된 안전 기준이 요구됩니다. 생성형 AI 챗봇 서비스를 출시하기 전 아동권리영향평가를 의무화하여 아동에게 정서적으로 미치는 영향이 있는지를 파악하고, 위험 대화 가능성은 없는지 판단해야 합니다. 즉 사후적 단순 필터링을 넘어, 챗봇의 대화 흐름 자체를 안전하게 설계하는 ‘Safe by Design’ 원칙으로 제도화될 필요가 있습니다.

02 생성형 AI 챗봇 이용 시 ‘대화 상대가 AI임’을 알 수 있도록 안내를 강화해야 합니다.

아동·청소년은 인공지능과 대화를 하다보면 상대가 사람인지 인공지능인지 혼동할 수 있습니다. 실제 초록우산의 조사에서도 이러한 경험이 확인되었습니다. 따라서 생성형 AI 챗봇은 사용자가 지금 대화하고 있는 대상이 인공지능임을 아동·청소년의 눈높이에 맞게 직관적인 표현으로 안내해야 합니다. 또한 자해나 폭력과 같은 위험한 대화내용이 등장할 경우, 시스템적으로 즉각 개입하는 표준화된 대응 프로토콜을 갖추어야 합니다.

03 생성형 AI 챗봇 관련 별도의 리터러시 교육이 필요합니다.

기존의 일반적인 AI 교육을 넘어, ‘챗봇과 대화하는 상황’ 자체를 전제로 한 리터러시 교육이 필요합니다. 초록우산의 실태조사에서도 다수의 아동은 본인의 나이, 이름 등 기본 정보를 생성형 AI 챗봇에 입력해 본 것으로 나타났습니다. 이는 이용 과정에서의 판단과 대응 역량을 강화하는 교육의 필요성을 보여주고 있습니다. 단순히 형식적인 교육이 아니라 생성형 AI 챗봇은 사람처럼 말하더라도 실제 감정과 의도가 없다는 점, 답변 내용은 정확하지 않을 수 있는 점, 위험한 대화 유도 및 답변 시 도움 요청 방법 등의 내용이 포함되어야 합니다. 특히 사례 기반의 체험형 교육을 확대하여, 아동·청소년이 다양한 상황을 직접 경험하고 판단해 보는 과정을 통해 인공지능의 대화 내용을 보다 비판적으로 수용할 수 있도록 해야 합니다. 이를 통해 인공지능 환경에서 스스로를 보호하고, 생성형 AI 챗봇을 ‘신뢰 대상’이 아닌 ‘도구’로서 AI 기술을 책임감 있게 활용할 수 있는 실질적 역량을 갖추도록 하는 것이 중요합니다.

04 생성형 AI 챗봇 중심의 아동권리 기반 정책 거버넌스 구축이 필요합니다.

생성형 AI 챗봇 정책 전반에 아동권리 관점을 체계적으로 반영하고 위험요소를 모니터링하는 거버넌스 구축이 필요합니다. 현행법상 인공지능사업자의 자율규제에 의존하는 것만으로는 아동·청소년이 생성형 AI 챗봇 서비스를 이용하며 발생하는 위험요소를 예방할 수 없습니다. 정부와 플랫폼은 자해·폭력·성적 유해 대화, 정서적 의존 유도 등 아동에게 중대한 영향을 미칠 수 있는 위험 유형을 표준화해야 하며, 일정 수준 이상의 위험이 탐지될 경우 사업자가 이를 분석·개선하고 그 결과를 공개하도록 하는 체계를 마련해야 합니다. 또한 과학기술정보통신부, 방송미디어통신위원회 등 관계 부처를 중심으로 아동·청소년 당사자, 보호자, 전문가가 참여하는 민관 협력 구조를 구축하여 생성형 AI 챗봇 관련 정책 수립 과정에서 다양한 이해관계자의 의견이 반영될 수 있도록 해야 합니다.

초록우산 온라인 세이프티 옹호활동

안전한 디지털 환경을 위한 초록우산의 노력

초록우산은 아동복지전문기관으로 모든 아동의 행복한 성장을 위해 돌봄·자립·교육·건강·안전·주거 등의 영역에서 복지사업을 강화해가고 있으며, 법·제도 및 인식개선 캠페인 등 다양한 옹호활동을 통해 아동의 권리가 보호·존중·실현되는 세상을 만들어가기 위해 노력하고 있습니다.

디지털 환경은 아동의 주요 생활 공간으로 자리잡았으나, 위험과 권리 침해에 대한 제도적 보호는 여전히 미흡한 상황입니다. 초록우산은 디지털 환경에서 발생하는 다양한 아동권리 침해에 대해 국가의 보호책임을 명확히 하고, 제도개선을 촉구하기 위한 옹호활동을 펼치고 있습니다.



인공지능과 아동권리 국제적 공동성명 동참



디지털 환경에서의 아동권리 보장을 위한 기자회견



디지털 역량 교육지원 조례 제안 활동



디지털 리터러시 교육지원

온라인세이프티 아카이브

초록우산이 지속적으로 추진해 온 온라인세이프티 옹호활동의 흐름을 한눈에 살펴볼 수 있는 공간입니다. 아동이 온라인에서 겪을 수 있는 위험을 줄이기 위한 활동과 연구를 바탕으로, 누구나 아동의 안전한 온라인 환경을 만들기 위한 참고 자료로 활용할 수 있도록 구성되었습니다.



아카이브 살펴보기

내 친구와의 위험한 대화

생성형 AI 챗봇 안전설계를 위한 이슈브리프 

발 행 일 2026년 4월

발 행 인 황 영 기

편 집 인 김 승 현

발 행 처 초록우산 옹호사업본부

주 소 서울시 중구 무교로 20

전화번호 1588-1940

